

Benchmark for Models Predicting Human Behavior in Gap Acceptance Scenarios

Julian F. Schumann, Jens Kober, Arkady Zgonnikov

Abstract—Autonomous vehicles currently suffer from a time-inefficient driving style caused by uncertainty about human behavior in traffic interactions. Accurate and reliable prediction models enabling more efficient trajectory planning could make autonomous vehicles more assertive in such interactions. However, the evaluation of such models is commonly oversimplistic, ignoring the asymmetric importance of prediction errors and the heterogeneity of the datasets used for testing. We examine the potential of recasting interactions between vehicles as gap acceptance scenarios and evaluating models in this structured environment. To that end, we develop a framework facilitating the evaluation of any model, by any metric, and in any scenario. We then apply this framework to state-of-the-art prediction models, which all show themselves to be unreliable in the most safety-critical situations.

Index Terms—autonomous vehicles, gap acceptance, behavior prediction, benchmark.

I. INTRODUCTION

SUCCESSFULLY implementing autonomous driving is one of the key technical challenges faced by the automotive industry as well as large parts of the research community, with tens of billions of dollars invested in recent years towards this goal [1]. The provision of those funds is motivated by several benefits promised by this technology. The foremost of these is safer driving, expressed by a significant decrease in accidents and, correspondingly, a reduction of bodily harm and financial losses. Additional advantages are also expected, such as more accessible mobility for people unable to drive or an easing of road congestion and traffic [2]–[4].

But despite all these investments, autonomous vehicles still suffer from many problems preventing widespread use [5], [6]. One such problem is their timidity in interactions with human traffic participants, caused by the uncertainty about the future behavior of those human agents. This uncertainty can prevent the autonomous vehicle from taking the most time-efficient actions if the resulting probability of a crash or near-crash is too high, resulting in the cautious driving style observed. Paradoxically, this can also be a safety risk, as such caution by an autonomous vehicle is often not expected by the surrounding humans, which can result in accidents such as being rear-ended [5], [7].

To reduce this uncertainty and to allow for a more efficient driving style without compromising on safety requirements,

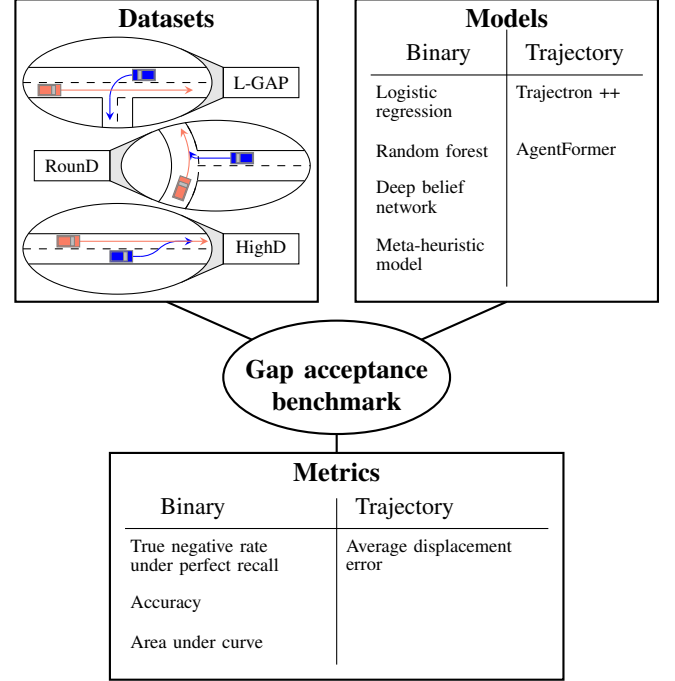


Fig. 1. The proposed framework allows researchers to evaluate the performance of any prediction model for human behavior according to any metric on any dataset including gap acceptance scenarios.

behavior prediction models can be used [8], [9], which project the future position of traffic participants. Those can range from models able to deal with any kind of traffic participant [10]–[12] to others being focused on predicting the behavior of a specific kind of participant, such as cars [13]–[19] or pedestrians [20]–[24].

However, the utility of those models—primarily designed to minimize the necessary trade-off between safety and efficiency in trajectory planning—is questionable, as the common methods for their evaluation diverge from the models’ purpose. First, most common metrics for evaluating prediction models, such as the final or average displacement error, ignore that the consequences of a false prediction are inherently asymmetric [25], [26]. For example, on a highway, wrong longitudinal predictions are far less dangerous than wrong lateral predictions, which might result in an autonomous vehicle reacting to a lane change too late. Second, the common approach of randomly selecting test cases from datasets [10], [12] is problematic due to the heterogeneity of those datasets, which typically include samples that can vary widely in their importance and difficulty. Such samples can range from a single vehicle following a

Manuscript received October 27, 2022; revised sometimes. (Corresponding Author: Julian Schumann)

The authors are with the Department of Cognitive Robotics, Delft University of Technology, Delft, Zuid Holland 2628 CD, The Netherlands (e-mail: j.f.schumann@tudelft.nl; j.kober@tudelft.nl; a.zgonnikov@tudelft.nl)

The source code, trained models, and data can be found online at [public Github repository](#)

lane to complex space-sharing conflicts with multiple agents at unsignalized intersections, where the behavior of human agents is often multi-modal and can change rapidly. Rare edge cases, where some traffic participants are very aggressive or even violate traffic rules and accidents are far more likely [27], [28], are also possible. But with randomly selected test cases, potentially poor performance in the most important situations can be compensated by good performance in less important but more numerous ones. For these reasons, a useless model might appear promising, which hampers further progress.

One possible approach to overcome these issues is including a path planning algorithm in evaluations, as suggested by Ivanovic and Pavone [25]. However, this adds further computational loads to an evaluation and only addresses the symmetry of common metrics, neglecting the varying difficulty and importance between testing samples. To cover both these problems, we suggest to instead narrow the evaluation to the most critical situations. In particular, we focus the evaluation of behavior prediction models on *gap acceptance* scenarios, a concept that encompasses most of the safety critical interactions between autonomous vehicles and humans [29]. In a gap acceptance scenario, an autonomous vehicle follows a particular trajectory over which a second traffic participant (e.g., a pedestrian or another vehicle) can move either in front of or behind the autonomous vehicle. Here, the first option (i.e., the human accepting the gap) would require the autonomous vehicle to potentially alter its trajectory planning, while the latter one of rejecting the gap would not. Due to the narrow focus, estimating the importance and difficulty of a particular situation can become much more straightforward. Additionally, as the human has only two options to decide between, such gap acceptance scenarios allow the usage of simple binary prediction models to estimate if the human behavior requires an adjustment of trajectory planning.

Many binary prediction models have been developed for gap acceptance scenarios. However, those mostly focus and are trained on a specific scenario, such as the street crossing behavior of pedestrians [30]–[32], the crossing behavior of cars at intersections [33]–[37], or lane change decisions on high ways [38]–[41]. Additionally, the development of those models still suffers from similar problems as the trajectory prediction models, such as the anisotropy of common metrics like accuracy. A random selection of test cases [31], [35] and neglect of the varying importance of different samples are also common. Additionally, in contrast to trajectory prediction models, which are commonly compared to each other on accepted benchmarks (such as on the ETH dataset [10]–[12] when predicting pedestrian crowds), an equivalent benchmark does not exist for binary prediction models [42]. Instead, those models are mostly trained and tested on datasets exclusive to the respective work and are—if at all—only compared against a small number of other selected models [30], [31], [35]–[37], [43]–[47].

Our goal in this work is to overcome these limitations of the current literature on both binary and trajectory prediction models and enable a meaningful evaluation of these models in gap acceptance scenarios. Such an evaluation cannot only make the development of trajectory prediction models more

goal-oriented, but also help determining to what extent the inclusion of specialized binary prediction models can improve the performance and reliability of general trajectory prediction models. To that end, this paper makes three main contributions:

- We develop a formal description of the gap acceptance process that applies to all possible gap acceptance scenarios. This description includes a detailed timeline of gap acceptance (Section II), which serves as a foundation for methods to estimate the criticality concerning the safety of each sample, which is a fundamental requirement for selecting meaningful test cases.
- We devise a framework for evaluating behavior prediction models in gap acceptance scenarios. This framework allows the integration of varied gap acceptance datasets, models, and evaluation metrics (Section III). It is inspired by similar works by Müller et al. for computer-based image retrieval algorithms [48], by Zaffar et al. in the field of visual place recognition [49], or Cao et al. for evaluating robustness to adversarial attacks of trajectory prediction models [50]. This approach would allow one to test any chosen model and compare it to other models in any possible environment more easily. Simultaneously, the framework allows precise control over the splitting the data into training and testing samples to evaluate the models’ reliability in the most difficult gap acceptance situations.
- Using the proposed framework, we compare several prediction models in their performance on three gap acceptance datasets, with a focus on their performance in safety-critical edge cases (Section IV and Figure 1). Consequently, other researchers can easily access the already implemented models to compare their own models against. Additionally, we compare models dedicated to predicting binary gap acceptance decisions to state-of-the-art trajectory prediction models to test the hypothesis that including dedicated binary models for gap acceptance problems could improve such trajectory prediction models.

II. DEFINING GAP ACCEPTANCE

To estimate the difficulty and control the importance of prediction task over disparate datasets, a coherent formal definition of gap acceptance scenarios is needed. Here we propose such a definition.

In a gap acceptance scenario, an autonomous vehicle V_E —also referred to as the ego-vehicle—plans to follow along a certain trajectory P_E along which it has the right of way. This trajectory overlaps with the trajectory P_T of another, human-controlled vehicle V_T (also named target vehicle). Such an overlap might, for example, happen at unsignalized intersections, where the agents move along crossing streets or on highways, where V_T wants to merge into the faster lane along which V_E is driving. In such situations, V_T can decide to move onto P_E either in front of or behind V_E , i.e., to accept or reject the gap offered by V_E . We assume that V_E has the right of way along P_E , as otherwise, traffic rules would obligate it to preemptively yield.

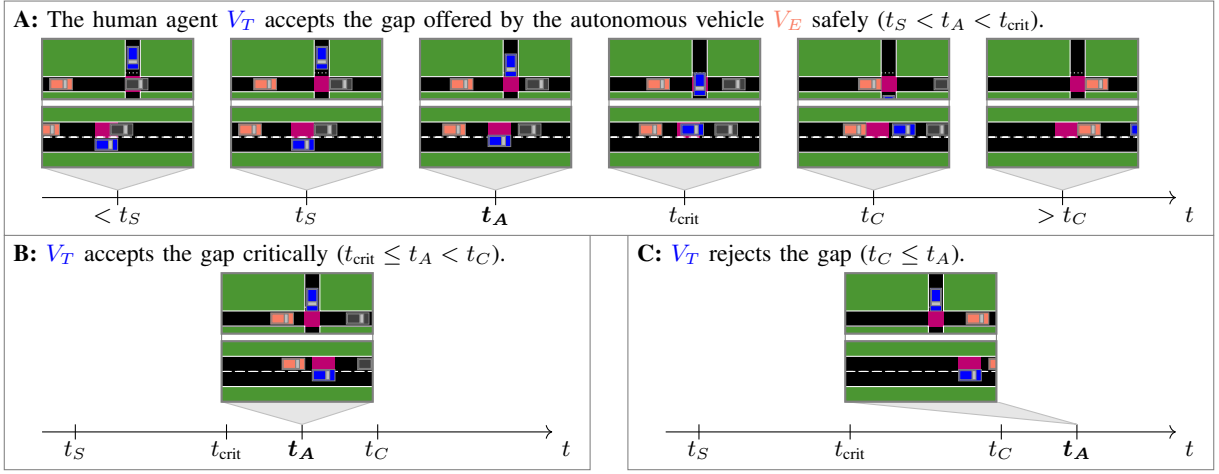


Fig. 2. The characteristic time-points of the gap acceptance process—defined by the relation of the agents to the contested space (in purple)—of two different examples of gap acceptance, intersection crossing (upper panels) and lane changing (lower panels). In both examples, the autonomous vehicle V_E (in red) offers a gap to the human driven vehicle V_T (in blue). In total, three cases are possible (A – C), depending on t_A , i.e., the time the target vehicle enters the contested space. In B, the accepted gap decision by the human is considered to be unsafe, as V_E cannot guarantee the avoidance of a crash, having potentially not enough time for braking. Meanwhile, in C, it might be possible that V_T crashes into V_E .

Under these conditions, a gap acceptance scenario is characterized by the spatio-temporal relation between the agents towards the so-called contested space [29]. There, the trajectories P_E and P_T would start to overlap, making this the location of a potential collision. An example is the overlap of two crossing lanes at an intersection. However, in specific scenarios (such as changing lanes on highways), the exact location of the meeting point of P_E and P_T can be at the discretion of the human agent V_T and therefore be unknown before the actual decision. In such cases, we then place the contested space under the assumption that V_T would decide to accept the gap immediately. For example, in the scenario of highway lane changes, the contested space would therefore move in parallel to V_T , only stopping to move once V_T starts to enter the lane of V_E .

The following time points then characterize the gap acceptance process (illustrated together with the contested space in Figure 2):

- t_S : At the starting time t_S , there is no longer any other vehicle along P_E in between V_E and the contested space. This is primarily the case when the vehicle preceding V_E leaves the contested space, but other options are imaginable, like the vehicle in front of V_E leaving P_E .
- t_C : At t_C , V_E starts to enter the contested space, closing the gap.
- $t_{\underline{C}}(t)$: A prediction of t_C by the ego vehicle, made at t , needed to allow gap size estimations during online applications. While this is scenario-dependent, the following condition has to be satisfied so that an open gap can still be characterized as such, even if V_E is moving away from the contested space:

$$\text{sgn}(t_{\underline{C}}(t) - t) = \text{sgn}(t_C - t) .$$

- t_{crit} : The last time V_E can safely prevent a collision even in the case of malicious behavior by V_T ; e.g., at this

point, a safe braking process could bring V_E to a stop before the intersection. t_{crit} can be formalized in the following condition:

$$\Delta t_D(t) = t_{\underline{C}}(t) - t - t_{brake}(t) = 0 \quad (1)$$

Here, the required braking time t_{brake} is not based on the maximum deceleration V_E is technically capable of, but instead, one that is considered safe. The time point t_{crit} is also the last time a prediction can be considered useful for further trajectory planning.

- t_A : At t_A , V_T enters the contested space, potentially accepting the gap.

We count V_T as rejecting the gap if V_E is allowed to move first onto the contested space, i.e., if $t_C \leq t_A$. If this is not the case and the human moves first ($t_A < t_C$), the gap is considered accepted.

III. FRAMEWORK FOR BENCHMARKING GAP ACCEPTANCE MODELS

After defining the fundamental characteristics of a gap acceptance scenario, we will use this groundwork to build a framework for benchmarking gap acceptance models. This framework should allow for the performance assessment of any model M on any dataset D according to any evaluation metric E . The following requirements need to be met for such an assessment to be possible and meaningful:

- R 1** The time point t_0 of a prediction must be controllable, as it influences not only the difficulty of the prediction but also its importance due to changing consequences of a false prediction.
- R 2** To evaluate models in critical situations, the framework should allow control over splitting all available samples into training and testing sets.
- R 3** Models producing (as well as metrics evaluating) binary or trajectory predictions should fit into the framework.

- Based on t_0 , n_I , n_O , and δt , the time-steps for input and output data are selected, named T_I and T_O respectively:

$$T_I = \{t_0 + i\delta t \mid i \in \{-n_I + 1, \dots, 0\}\}$$

$$T_O = \{t_0 + i\delta t \mid i \in \{1, \dots, n_O\}\}$$

- For those time-steps, the input trajectories X_{T_I} and output trajectories X_{T_O} are extracted from X_T , using interpolation if necessary. For the output, only the trajectory of the target agent V_T is required, as its behavior is to be predicted.
- Certain domain information k is collected, for instance, the location at which the trajectories X_T were collected or the test subjects involved in gathering the data.

The input data D_I then includes from each sample the input trajectory X_{T_I} and the corresponding time-steps T_I . Meanwhile, the output data D_O takes the output trajectory X_{T_O} and the corresponding time-steps T_O , as well as the binary decision a , the time of accepting the gap t_A , and the domain information k from each sample. More mathematical details are in Appendix A-C.

C. Creating training and testing set — Splitting method S

To fulfill requirement R 2, the splitting method S , separating the given samples into training and testing sets, is a crucial part of the framework. As this functionality should be independent of the scenario, it is part of the separate splitting module. Examples for this range from random splitting to methods taking into account all the information in D_I and D_O .

Besides the potential similarity measure ζ of training samples to the training set, this functionality creates the training data $D_{I,\text{train}}$ and $D_{O,\text{train}}$ as well as the test data $D_{I,\text{test}}$ and $D_{O,\text{test}}$.

D. Training the model on the training set — Model M

After splitting the samples into training and testing sets, the model has to be trained, which is one of the functionalities of the model module that has to be individually implemented for each model. If the model, for instance, requires input velocities, extracting those from the given position data X_{T_I} and X_{T_O} is done here.

E. Making the predictions for the testing set — Model M

For every sample from the input testing set $D_{I,\text{test}}$, a prediction d_{pred} is made by the trained model, with all d_{pred} constituting the set of predictions $D_{OP,\text{test}}$. As such predictions rely on a trained model, this functionality is also part of the model module. Depending on this model, each prediction d_{pred} might take different forms. Three different forms of stochastic predictions are implemented into the framework.

- Binary prediction: $d_{\text{pred}} = a_{\text{pred}}$, i.e., only the probability $a_{\text{pred}} \in [0, 1]$ of V_T accepting the offered gap is predicted.
- Timing prediction: $d_{\text{pred}} = \{a_{\text{pred}}, t_{A,\text{pred}}\}$, i.e., not only a_{pred} is predicted but also the time $t_{A,\text{pred}}$ at which the gap might be accepted.
- Trajectory prediction: $d_{\text{pred}} = X_{T_O,\text{pred}}$, i.e., the full trajectory of V_T is predicted. This prediction consists of

n_p trajectories $X_{T_O,p}$ —all equally likely—to represent probabilistic outputs.

More mathematical details are in Appendix A-D.

F. Transforming the predictions — Dataset D

To fulfill requirement R 3, we then need to be able to transform a prediction d_{pred} to any other form we might require. This functionality is part of a specific dataset, as this entails the required context information to, for example, classify different trajectories.

These transformations rely on three different functions T_i :

- T_1 : takes the trajectory prediction $X_{T_O,\text{pred}}$ and then provides $\{a_{\text{pred}}, t_{A,\text{pred}}\}$. This is similar to the extraction of the time points in Section III-B.
- T_2 : takes the prediction $\{a_{\text{pred}}, t_{A,\text{pred}}\}$ and then provides the trajectory prediction $X_{T_O,\text{pred}}$, consisting of n_p trajectories from the predictions of two conditional trajectory prediction models trained only on accepted and rejected gaps respectively. These are selected so that $T_1(X_{T_O,\text{pred}})$ results in the original inputs.
- T_3 : takes a binary prediction a_{pred} and provides the predicted time of accepting the gap $t_{A,\text{pred}}$, by extracting it from the prediction of a trajectory prediction model trained only on accepted gaps.

These three functions (exact implementation in Appendix A-E) are enough to enable any transformation between prediction forms, as seen in Table I.

G. Evaluating the predictions — Metric E

This functionality—the main part of the metric module—implements the performance evaluation, comparing the actual outputs $D_{O,\text{test}}$ with the predicted outputs $D_{OP,\text{test}}$. It returns either a combined value F or instead a separate value for each sample $d_{\text{pred}} \in D_{OP,\text{test}}$, resulting in the output F .

IV. BENCHMARK IMPLEMENTATION

We implemented the framework described above by linking together several datasets, models, splitting methods, and metrics. These were chosen not to comprehensively cover all possible gap acceptance scenarios and prediction models but to demonstrate the flexibility and utility of the proposed framework. Still, our implementation can already serve as a benchmark for new prediction models. This section only presents an overview of the implementation, with full technical specification provided in [supplementary materials](#).

A. Datasets

Different datasets are implemented into the framework (Table II), including data recorded on real roads as well as data from a driving simulator study. The naturalistic datasets used here are captured by drones and distinguished by accurate position labeling. They cover lane changes on German highways (the *highD* dataset [51]) and roundabouts (the *round* dataset [52]). The L-GAP dataset covers left turns at unsignalized intersections through oncoming traffic recorded in a driving simulator [35]. It has been chosen due to the simplicity of its environment, contrasting the more complex scenarios in the naturalistic datasets.

TABLE I
TRANSFORMATION BETWEEN ANY POSSIBLE PREDICTION TYPES d_{pred} USING THE THREE IMPLEMENTED FUNCTIONS T_i .

Input $\Rightarrow$ Output	a_{pred}	$\{a_{\text{pred}}, t_{A,\text{pred}}\}$	$\mathbf{X}_{T_O,\text{pred}}$
Binary prediction a_{pred}	—	$\{a_{\text{pred}}, T_3(a_{\text{pred}})\}$	$T_2(\{a_{\text{pred}}, T_3(a_{\text{pred}})\})$
Timing prediction $\{a_{\text{pred}}, t_{A,\text{pred}}\}$	—	—	$T_2(\{a_{\text{pred}}, t_{A,\text{pred}}\})$
Trajectory prediction $\mathbf{X}_{T_O,\text{pred}}$	in $T_1(\mathbf{X}_{T_O,\text{pred}})$	$T_1(\mathbf{X}_{T_O,\text{pred}})$	—

TABLE II
THE NUMBER OF ACCEPTED GAPS N_A AND REJECTED GAPS N_{-A} IN THE IMPLEMENTED DATASETS, IN THE FORM: $N_A - N_{-A}$ (MEDIAN $t_{\underline{C}}(t_0)$). THE NUMBERS DEPEND ON THE METHOD FOR CHOOSING THE TIME OF PREDICTION t_0

Dataset	Initial gaps at their opening	Gaps with fixed size	Critical gaps
highD (Lane changes)	1406 – 7026 (8.9 s)	461 – 1568 (11.8 s)	0 – 7025 (0.8 s)
highD (Lane changes - restricted)	1406 – 1001 (12.9 s)	392 – 241 (8.7 s)	0 – 1000 (0.7 s)
rounD (Roundabout)	662 – 917 (2.5 s)	168 – 168 (2.9 s)	33 – 913 (1.0 s)
L-GAP (Left turns)	703 – 724 (4.6 s)	496 – 572 (3.5 s)	369 – 723 (2.3 s)

1) *Lane changes*: Here we focus on lane changes of the target vehicle V_T toward a faster lane to the left, along which the ego vehicle V_E driving there has the right of way. While it could be argued that predictions in such situations could be simply based on turn signals, one cannot rely on human drivers to correctly use these [47]. As a source of lane change data, we used two versions of the *highD* dataset, *full* and *restricted*.

The full *highD* dataset, not employing any filters, is heavily biased toward trajectories without a lane change. This is not a problem per se, but in such trajectories it is not known whether the target vehicle V_T even had an intention to change lanes (i.e., if there was a gap acceptance situation in the first place). For this reason, in addition to the full *highD* dataset, we added a restricted version of it which only included samples for which it can be inferred that the target vehicle V_T indeed considered a lane change. Criteria are either a lane change of V_T after V_E has passed or V_T braking to not collide with the preceding vehicle instead of changing lanes. Still, in both versions of *highD*, the gaps are always accepted with large safety margins (Table II).

2) *Roundabout*: In the *rounD* dataset, the target vehicle V_T has to enter a roundabout, which it can do in front of or behind the ego vehicle V_E already in the roundabout. As the trajectories are recorded in Germany, the ego vehicle inside the roundabout has the right of way. Compared to *highD*, this dataset is far more balanced between accepted and rejected gaps, but still only includes few critically accepted gaps.

3) *Left turns*: In the *L-GAP* dataset [35], the driver of the target vehicle V_T intends to turn left at an intersection. The driver had to decide whether to do this in front of or behind the ego vehicle V_E approaching the intersection from the opposite direction with the right of way. While the number of samples in this dataset is comparatively small, they are relatively balanced between accepted and rejected gaps. Also, they include many gaps accepted after t_{crit} (Table II). Nonetheless, as V_T starts in an idling position at some distance to the contested area, this might not be the most challenging dataset, as an onset of movement before t_0 in most cases is

an apparent indicator of V_T intending to accept the gap.

B. Test-train Splitting Methods

Two splitting methods are implemented, without a method for calculating the similarity measure ζ . Nonetheless, to enable at least a qualitative approximation of a model's robustness, the methods are designed to produce testing sets of varying difficulty for the prediction models.

The easier variant performs a stratified random splitting, while the second, more extreme method sorts the most unintuitive behavior of the target vehicle into the testing set (e.g. accepting a very small gap or rejecting a very large gap).

In both cases, the testing set includes 20% of the samples and the training set the remaining 80%.

C. Models

The benchmark includes two state-of-the-art trajectory prediction models

- *Trajectron++* (also referred to as $T+$), a deep-learning model mainly based on long-short-term memory cells [10].
- *AgentFormer (AF)*, a deep-learning model based on transformers [12]. Compared to $T+$, it has ten times more trainable parameters.

For the binary prediction models for gap acceptance, there is, as mentioned above, a lack of a common benchmark, making the models' selection more contentious. Four models have been selected nonetheless:

- Logistic regression (*LR*) is commonly used for predicting human gap acceptance decisions [32] and is therefore included as a simple baseline.
- Random forests (*RF*) have been shown to outperform other approaches such as logistic regression and standard decision trees in gap acceptance prediction [33].
- Deep belief networks (*DB*), also used previously to predict human gap acceptance decisions [39].

- A metaheuristic model based on combining all other five models above (*MH*); previously a similar approach for lane changes has been shown to outperform each of the models included in it [41].

The benchmark does not include any dedicated timing prediction models yet, as their primary representative, the drift-diffusion model [31], [35], can currently not be trained on datasets with a large number of unique samples in a reasonable amount of time. Nonetheless, to allow for future expansion of the benchmark, the framework has been designed with such models in mind.

D. Evaluation Metrics

We have included several metrics that characterize models in terms of quality of binary predictions (accept/reject gap) as well as full trajectory predictions. The following metrics are commonly used in the literature:

- **Accuracy:** This metric is a widespread method to evaluate the performance for binary prediction models [33], [39], [41]. However, accuracy is a symmetric metric, i.e. it is unable to differentiate between false negative and false positive predictions and is best used in cases where $t_0 \ll t_{\text{crit}}$, as the consequences of false predictions are not too different there (Appendix A-A).
- **AUC:** This metric for binary prediction models, the Area Under Curve of a receiver operating characteristics curve, addresses one central point of criticism of the accuracy metric, namely its sensitivity to biases in the testing set. Nonetheless, like the accuracy metric, it does not consider the potentially differing severity of false predictions. Hence, we only apply it to rate a prediction model's performance when $t_0 \ll t_{\text{crit}}$.
- **ADE_β :** This metric, the Average Displacement Error of the $n_p\beta$ least erroneous predicted trajectories, is commonly applied to trajectory predictions [10]–[12], with $\beta = 1$ or $\beta = 0.05$ being used in this work. As this metric also does not take into account the severity of different false predictions [25], [26] and requires equally long prediction horizons as well, it is only applied for constant gap sizes ($t_0 = \min\{t \mid t_{\text{C}}(t) - t = \Delta t\} \ll t_{\text{crit}}$).

In addition to the above metrics, we also propose a novel metric that takes into account potential consequences of a wrong prediction.

- **TNPR:** The True Negative rate under Perfect Recall is applied to binary predictions, where the threshold for categorizing a prediction a_{pred} as an accepted gap is set so low that there are no false negative predictions (perfect recall). It considers the different consequences of false negative and positive predictions, assuming that collisions are to be avoided at all costs. It consequently estimates the likelihood of not having to brake needlessly for rejected gaps while guaranteeing safe interactions. This metric is designed specifically for critical gap size with $t_0 \approx t_{\text{crit}}$, as at earlier time points the need for a perfect recall would be unreasonable.

V. RESULTS

Our benchmark provides insights into performance of tested models under different conditions (Figures 4, 5). There, when provided with more input time steps ($n_I = 10$ vs. $n_I = 2$), i.e., more information to extract signs of future behavior from, the models' predictions were generally better. These observations are as expected as the worsening performance of models tested on the most unintuitive samples from the extreme splitting case. There, the models have to extrapolate, a typically far more difficult task than the interpolation needed for predicting random samples similar to the training set [53]. The poor performance on unintuitive samples is especially pronounced when looking at the *TNPR* at critical gaps (Figure 5), where no model could substantially outperform a random predictor on both datasets.

Furthermore, it can be seen in Figure 4 that accuracy, as it depends on the distribution of accepted and rejected gaps in the test set, is not a suitable metric for comparing model's performance on different datasets and different prediction times. The same can be said of the *ADE*. Instead, *AUC* seems a far more reliable metric, being the only metric that accurately captures that most models have fewer problems with the restricted lane change scenario than the full one.

When comparing the difficulty of different scenarios, it can be seen that, generally, the prediction of human behavior at roundabouts seems to be easier than predicting lane change behavior. For those two scenarios, it can also be observed that predictions made for constant gap sizes seem slightly more accurate than those made for initial gaps. The only difference here is the left turn scenario, where this difference is far more pronounced, especially for $n_I = 2$. We can explain those difficulties by the initial lack of motion in the target vehicle V_T . The difference is lesser for $n_I = 10$, as the prediction is made for this scenario at $t_0 = t_S + (n_I - 1) * \delta t$ (instead of $t_0 = t_S$) due to the lack of trajectory data before t_S . Consequently, for $n_I = 10$, an onset of motion is far more likely to be recorded than for $n_I = 2$. Overall, predictions made for constant time gaps on the left turn scenario are relatively easy, while they are relatively hard when made at the opening of the gap.

Besides the general trends mentioned above, which are mostly valid for all models, we can also compare those models among themselves. When comparing the trajectory prediction models, it can be seen that the Trajectron++ model (*T+*) consistently outperforms the AgentFormer model (*AF*). This is surprising, as *AF* previously outperformed *T+* on pedestrian trajectory prediction benchmarks [12]. This contradiction might be explained by over-fitting the many trainable parameters for *AF* on relatively small datasets here. Meanwhile, the logistic regression (*LR*) model is often the most promising approach for binary prediction models, especially when tested on random samples. Together, those results indicate that increasing the complexity of such models and their number of parameters might not be a panacea, with simpler models being more promising, especially if datasets are relatively small.

When we compare binary models against trajectory prediction models, we can observe differing behavior for different

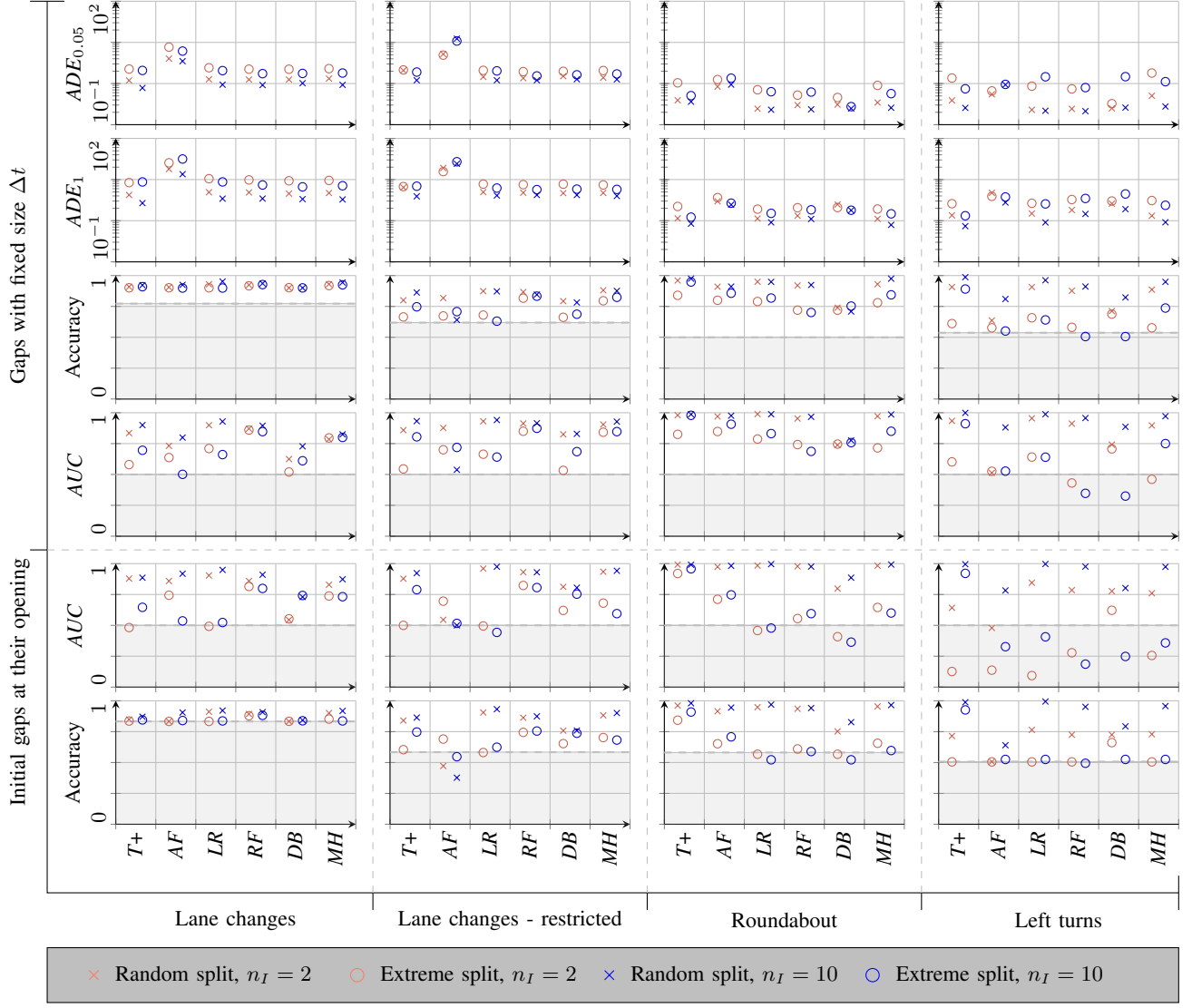


Fig. 4. The results of evaluating behavior prediction models in gap acceptance scenarios on different datasets, with prediction being made either at the initial opening of the gap ($t_0 = t_S$) or with fixed gap sizes ($t_0 = \min\{t | t_{\bar{C}}(t) - t = \Delta t\}$). The color indicates the number of input time-steps n_I given to the models, and the marker type denotes the splitting method, which can be *random* or *extreme*. The dashed gray lines indicate the performance F_r of a uniformly random binary predictor.

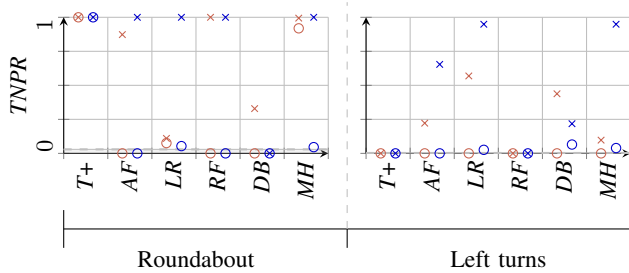


Fig. 5. The true negative rate under perfect recall (TNPR) of different prediction models, using the same visualizations as in Figure 4, tested on the last useful predictions ($t_0 = t_{\text{crit}} - t_\epsilon$).

scenarios. On the one hand, the best performance on the lane change datasets is generally achieved by a binary prediction model, while on the other hand, the reverse is the case on the

other two datasets. One main difference here is the prediction horizon ($\Delta t \approx 10$ s and $\Delta t < 4$ s respectively), which might explain those results. That the average displacement errors are much more noticeable in the lane change scenario also supports this explanation, further showing the difficulties of using trajectory prediction under those conditions.

Lastly, when evaluating the promise of conditional trajectory prediction models, we can hold that the benefits are negligible in most cases, except at roundabout and left turns for $ADE_{0.05}$. Nonetheless, due to the problems with that metric, more than those results are needed to render a final judgment. Similarly, due to the small size of datasets and the low number of models, the previous results should also be treated carefully.

VI. CONCLUSION

We proposed a framework that connects previously disparate datasets, models, and metrics in the benchmark for testing

behavior prediction models in gap acceptance scenarios. We demonstrated its potential by comparing two state-of-the-art trajectory prediction models with several binary gap acceptance models. Additionally, we showed that relying on the characteristic time points of gap acceptance scenario to select the most unintuitive samples in the splitting module is a promising approach to analyzing model generalization, as seen by the general decrease in performance for models trained on those samples. Our framework is open-source and specifically designed in a way to simplify adding new datasets, splitting methods, models, and evaluation metrics, which allows researchers to expand it in future.

One particularly important addition to the benchmark would be a datasets containing more critically accepted gaps, as this would allow for an increased meaningfulness of metrics applied to last useful predictions. Additionally, a metric better aligned with the main purpose of a prediction model as a part of an autonomous vehicle is still needed. Likewise, currently there exists no method for calculating the similarity between testing and training set ζ ; the future addition of this would permit a quantitative comparison of a model's robustness against unintuitive test samples. Lastly, one could expand the framework to provide scenario-independent inputs similar to t_C , which would make training a model on two unrelated datasets simultaneously possible, leading to a better estimation of the models generalizability.

We acknowledge that testing a model in gap acceptance scenarios alone is necessary, but not sufficient for justifying its usage in actual vehicles. Consequently, expanding the framework to non-gap-acceptance scenarios is an important avenue for future research. This will enable a more holistic testing but only for models predicting (and metrics evaluating) trajectories. Nonetheless, performance of models on non-gap acceptance scenarios should still be given lower priority compared to gap acceptance scenarios which are more safety-critical.

Our results resonate with the recent literature on hybrid AI [54], [55], showing that including binary prediction models in specific scenarios might make data-driven trajectory prediction models more reliable, especially in accurately predicting dangerous situations. However, especially for the unintuitive and safety-critical edge cases, most models often performed only slightly better than a random predictor at best. Therefore, there currently seems to be no model that a trajectory planning algorithm in any scenario can rely on to substantially increase the effectiveness of an autonomous vehicle's driving style, necessitating further research into such models.

APPENDIX A

DETAILING THE FRAMEWORK FOR BENCHMARKING GAP ACCEPTANCE MODELS

A. Evaluating the predictions - Metric E

Due to the differing consequences of false negative and false positive predictions when using binary predictions a_{pred} , there are limitations on which metrics are usable at certain t_0 . For $t_0 \ll t_{\text{crit}}$, the consequences of a wrong prediction are generally minor, as time is left to wait for future information before

more significant changes to trajectory planning are necessary. Furthermore, even if the target vehicle would immediately accept the gap after t_0 , the necessary response is likely neither uncomfortable nor risky. Consequently, symmetric metrics can be used here.

For $t_0 \approx t_{\text{crit}}$ however, no more time for further observations is left, resulting in far more severe consequences for both false positive and false negative predictions. A false positive prediction would unnecessarily result in a harsh and uncomfortable braking maneuver. Meanwhile, a wrong negative prediction leads to an unsafe gap acceptance maneuver, with the safety of the interaction between V_E and V_T no longer in the control of the autonomous vehicle V_E . Accidents or the need for dangerous emergency maneuvers, which could result in material damage or even bodily harm, are then possible. As the latter should be avoided at all costs, a false negative prediction at this time is far more consequential, which should be reflected in the evaluation metric.

B. Setting the prediction time - Metric E

The method for determining the time t_0 can impact the size of the resulting dataset (Table II) due to the condition from Equation (2). Namely, for constant gap size ($t_0 = \min\{t | t_C(t) - t = \Delta t\}$), the number of available samples will be reduced, as all gaps with an initial smaller gap size ($t_C(t_S) - t_S < \Delta t$) will be excluded. The same is the case for gaps already accepted before t_0 . For critical gap sizes instead (i.e., $t_0 = t_{\text{crit}} - t_\epsilon$), all gaps accepted before t_{crit} will be excluded, leading to extremely biased datasets, sometimes even removing all accepted gaps.

C. Extracting input and output - Dataset D

Trajectories $\mathbf{X}_T$ are only provided over the interval $I_T = [\max \mathbf{T}, \min \mathbf{T}]$. The conditions C_S , C_C , and C_A are then used to extract t_S , t_C , and t_A respectively, with $\mathbf{T}_{C_i} = \{t | C_i(t) \forall t \in I_T\}$:

$$\begin{aligned} t_S = T_S(\mathbf{X}_T) &= \begin{cases} \min \mathbf{T} & \mathbf{T}_{C_S} = \emptyset \\ \max \mathbf{T}_{C_S} & \text{else} \end{cases} \\ t_C = T_C(\mathbf{X}_T) &= \begin{cases} t_C(\max \mathbf{T}) & \mathbf{T}_{C_C} = \emptyset \\ \min \mathbf{T}_{C_C} & \text{else} \end{cases} \\ t_A = T_A(\mathbf{X}_T) &= \begin{cases} \max \mathbf{T} + t_\epsilon & \mathbf{T}_{C_A} = \emptyset \\ \min \mathbf{T}_{C_A} & \text{else} \end{cases} \end{aligned} \quad (\text{A.3})$$

If $\mathbf{T}_{C_C} = \emptyset \wedge \mathbf{T}_{C_A} = \emptyset$, the sample will be excluded from the dataset, as no decision can be observed. If not, the number of output time steps n_O is calculated by the following equation:

$$n_O = \left\lceil \frac{t_C - t_0}{\delta t} \right\rceil \quad (\text{A.4})$$

This is chosen so that $T_A(\mathbf{X}_{T_O}) \leq \max \mathbf{T}_O$ is sufficient to categorize the gap as accepted.

D. Making the predictions for the testing set - Model M

The predicted time $t_{A,\text{pred}}$ are the decile values of the underlying probability distribution represented by the set T_A :

$$t_{A,\text{pred}} = \{Q_{t_A}(p) \mid p \in \{0.1, 0.2, \dots, 0.9\}\} = Q_9(T_A) \in \mathbb{R}^9 \quad (\text{A.5})$$

Here, Q_{t_A} is the quantile function associated with this underlying distribution of t_A . Meanwhile, the predicted stochastic trajectory $X_{T_O,\text{pred}}$ is expressed by using n_p deterministic trajectories $X_{T_O,p}$:

$$X_{T_O,\text{pred}} = \{X_{T_O,1}, \dots, X_{T_O,n_p}\} \quad (\text{A.6})$$

E. Transforming the predictions - Dataset D

When implementing the transformation of a prediction d_{pred} into another form, two instances of the trajectory prediction model *Trajectron++* [10] are used, namely M_A , trained on all samples from the specific dataset D (D_I and D_O) where $a = 1$, and $M_{\neg A}$, trained on the remaining samples with $a = 0$. Furthermore, the function f_a is defined:

$$f_a(X_{T_O,p}) = \begin{cases} 1 & T_A(X_{T_O,p}) \leq \max T_O \\ 0 & \text{else} \end{cases} \quad (\text{A.7})$$

T_1 : The use of f_a results in

$$a_{\text{pred}} = \frac{1}{n_p} \sum_p f_a(X_{T_O,p})$$

$$t_{A,\text{pred}} = Q_9(\{T_A(X_{T_O,p}) \mid f_a(X_{T_O,p}) = 1\}) \quad (\text{A.8})$$

T_2 : M_A and $M_{\neg A}$ are used to respectively predict $X_{T_O,\text{pred},A}$ and $X_{T_O,\text{pred},\neg A}$, resulting in

$$X_{T_O,\text{pred}} = \{X_{T_O,p,\neg A} \mid p \in R_{\neg A}\} \cap \{X_{T_O,p,A} \mid p \in R_A\}, \quad (\text{A.9})$$

where

$$R_{\neg A} = R(n_p(1 - a_{\text{pred}}), \{p \mid f_a(X_{T_O,p,\neg A}) = 0\}, \mathbf{1})$$

$$R_A = R(n_p a_{\text{pred}}, \{p \mid f_a(X_{T_O,p,A}) = 1\}, \mathbf{W}_A). \quad (\text{A.10})$$

Here, $R(m, M, W)$ is a function that randomly selects m samples from M under weight W . W_A is chosen so that the sum of all weights of all $T_A(X_{T_O,p,A})$ in each quantile of $t_{A,\text{pred}}$ is identical. This approach of using conditional trajectory prediction models is inspired by Xie et al. [39].

T_3 : Here, one uses M_A to get $X_{T_O,\text{pred},A}$, resulting in

$$\{a_{\text{pred},A}, t_{A,\text{pred}}\} = T_1(X_{T_O,\text{pred},A}) \quad (\text{A.11})$$

In all cases, $p \in \{1, \dots, n_p\}$ can be assumed.

REFERENCES

- [1] D. Holland-Letz, M. Kässer, B. Kloss, and T. Müller, "Mobility's future: An investment reality check | McKinsey," Apr. 2021.
- [2] J. S. Brar and B. Caulfield, "Impact of autonomous vehicles on pedestrians' safety," in *IEEE Int. Conf. on Intell. Transp. Syst. (ITSC)*, pp. 714–719, Oct. 2017.
- [3] J. Meyer, H. Becker, P. M. Bösch, and K. W. Axhausen, "Autonomous vehicles: The next jump in accessibilities?," *Research Transp. Econ.*, vol. 62, pp. 80–91, June 2017.
- [4] J. Pisarov and G. Mester, "The future of autonomous vehicles," *FME Trans.*, vol. 49, no. 1, pp. 29–35, 2021.
- [5] M. Milford, S. Anthony, and W. Scheirer, "Self-Driving Vehicles: Key Technical Challenges and Progress Off the Road," *IEEE Potentials*, vol. 39, pp. 37–45, Jan. 2020.
- [6] J. Wang, L. Zhang, Y. Huang, and J. Zhao, "Safety of Autonomous Vehicles," *J. Adv. Transp.*, vol. 2020, p. e8867757, Oct. 2020.
- [7] A. Sinha, S. Chand, V. Vu, H. Chen, and V. Dixit, "Crash and disengagement data of autonomous vehicles on public roads in California," *Sci. Data*, vol. 8, p. 298, Dec. 2021.
- [8] D. Sadigh, S. Sastry, S. A. Seshia, and A. D. Dragan, "Planning for Autonomous Cars that Leverage Effects on Human Actions," in *Robotics: Sci. Syst. XII*, 2016.
- [9] S. Mozaffari, O. Y. Al-Jarrah, M. Dianati, P. Jennings, and A. Mouzakitis, "Deep Learning-Based Vehicle Behavior Prediction for Autonomous Driving Applications: A Review," *IEEE Trans. on Intell. Transp. Syst.*, vol. 23, pp. 33–47, Jan. 2022.
- [10] T. Salzmann, B. Ivanovic, P. Chakravarty, and M. Pavone, "Trajectron++: Dynamically-Feasible Trajectory Forecasting with Heterogeneous Data," in *Comput. Vis. – ECCV 2020* (A. Vedaldi, H. Bischof, T. Brox, and J.-M. Frahm, eds.), Lecture Notes in Computer Science, pp. 683–700, 2020.
- [11] F. Giuliari, I. Hasan, M. Cristani, and F. Galasso, "Transformer Networks for Trajectory Forecasting," in *Int. Conf. on Pattern Recognit. (ICPR)*, pp. 10335–10342, Jan. 2021.
- [12] Y. Yuan, X. Weng, Y. Ou, and K. M. Kitani, "AgentFormer: Agent-Aware Transformers for Socio-Temporal Multi-Agent Forecasting," in *IEEE/CVF Int. Conf. on Comput. Vis.*, pp. 9813–9823, 2021.
- [13] F. Diehl, T. Brunner, M. T. Le, and A. Knoll, "Graph Neural Networks for Modelling Traffic Participant Interaction," in *IEEE Intell. Veh. Symp. (IV)*, pp. 695–701, June 2019.
- [14] S. Kolekar, J. de Winter, and D. Abbink, "Human-like driving behaviour emerges from a risk-based driver model," *Nat. Commun.*, vol. 11, p. 4850, Dec. 2020.
- [15] X. Huang, G. Rosman, I. Gilitschenski, A. Jasour, S. G. McGill, J. J. Leonard, and B. C. Williams, "HYPER: Learned Hybrid Trajectory Prediction via Factored Inference and Adaptive Sampling," in *IEEE Int. Conf. on Robotics Autom. (ICRA)*, pp. 2906–2912, May 2022.
- [16] A. Cui, S. Casas, A. Sadat, R. Liao, and R. Urtasun, "LookOut: Diverse Multi-Future Prediction and Planning for Self-Driving," in *IEEE/CVF Int. Conf. on Comput. Vis.*, pp. 16107–16116, 2021.
- [17] R. Chandra, A. Bera, and D. Manocha, "Using Graph-Theoretic Machine Learning to Predict Human Driver Behavior," *IEEE Trans. on Intell. Transp. Syst.*, vol. 23, pp. 2572–2585, Mar. 2022.
- [18] L. Sun, X. Jia, and A. Dragan, "On Complementing End-To-End Human Behavior Predictors with Planning," *Robotics science systems*, Jan. 2021.
- [19] Z. Cao, E. Biyik, G. Rosman, and D. Sadigh, "Leveraging Smooth Attention Prior for Multi-Agent Trajectory Prediction," in *IEEE Int. Conf. on Robotics Autom. (ICRA)*, pp. 10723–10730, May 2022.
- [20] A. Bighashdel, P. Meletis, and G. Dubbelman, "Towards Equilibrium-based Interaction Modeling for Pedestrian Path Prediction," in *IEEE Int. Conf. on Intell. Transp. Syst. (ITSC)*, pp. 1–8, Sept. 2020.
- [21] A. Sadeghian, V. Kosaraju, A. Sadeghian, N. Hirose, H. Rezaatoughi, and S. Savarese, "SoPhie: An Attentive GAN for Predicting Paths Compliant to Social and Physical Constraints," in *IEEE/CVF Conf. on Comput. Vis. Pattern Recognit. (CVPR)*, pp. 1349–1358, June 2019.
- [22] P. Kothari, B. Siffringer, and A. Alahi, "Interpretable Social Anchors for Human Trajectory Forecasting in Crowds," in *IEEE/CVF Conf. on Comput. Vis. Pattern Recognit. (CVPR)*, pp. 15551–15561, June 2021.
- [23] F. Camara, N. Bellotto, S. Cosar, F. Weber, D. Nathanael, M. Althoff, J. Wu, J. Ruenz, A. Dietrich, G. Markkula, A. Schieben, F. Tango, N. Merat, and C. Fox, "Pedestrian Models for Autonomous Driving Part II: High-Level Models of Human Behavior," *IEEE Trans. on Intell. Transp. Syst.*, vol. 22, pp. 5453–5472, Sept. 2021.
- [24] J. Yue, D. Manocha, and H. Wang, "Human Trajectory Prediction via Neural Social Physics," in *Lect. Notes Comput. Sci.*, July 2022.
- [25] B. Ivanovic and M. Pavone, "Injecting Planning-Awareness into Prediction and Detection Evaluation," in *IEEE Intell. Veh. Symp. (IV)*, pp. 821–828, June 2022.
- [26] A. Farid, S. Veer, B. Ivanovic, K. Leung, and M. Pavone, "Task-Relevant Failure Detection for Trajectory Predictors in Autonomous Vehicles," July 2022. arXiv:2207.12380 [cs].
- [27] S. Noh, "Probabilistic Collision Threat Assessment for Autonomous Driving at Road Intersections Inclusive of Vehicles in Violation of Traffic

- Rules,” in *IEEE/RSJ Int. Conf. on Intell. Robots Syst. (IROS)*, pp. 4499–4506, Oct. 2018.
- [28] Q. Zhang, S. Hu, J. Sun, Q. A. Chen, and Z. M. Mao, “On Adversarial Robustness of Trajectory Prediction for Autonomous Vehicles,” in *IEEE/CVF Conf. on Comput. Vis. Pattern Recognit.*, pp. 15159–15168, 2022.
- [29] G. Markkula, R. Madigan, D. Nathanael, E. Portouli, Y. M. Lee, A. Dietrich, J. Billington, A. Schieben, and N. Merat, “Defining interactions: a conceptual framework for understanding interactive behaviour in human and automated road traffic,” *Theor. Issues Ergon. Sci.*, vol. 21, Feb. 2020.
- [30] S. K. Jayaraman, L. P. Robert, X. J. Yang, and D. M. Tilbury, “Multi-modal Hybrid Pedestrian: A Hybrid Automaton Model of Urban Pedestrian Behavior for Automated Driving Applications,” *IEEE Access*, vol. 9, pp. 27708–27722, 2021.
- [31] J. Pekkanen, O. T. Giles, Y. M. Lee, R. Madigan, T. Daimon, N. Merat, and G. Markkula, “Variable-Drift Diffusion Models of Pedestrian Road-Crossing Decisions,” *Comput. Brain & Behavior*, vol. 5, pp. 60–80, Mar. 2022.
- [32] A. Theofilatos, A. Ziakopoulos, O. Oviedo-Trespalacios, and A. Timmis, “To cross or not to cross? Review and meta-analysis of pedestrian gap acceptance decisions at midblock street crossings,” *J. Transport & Health*, vol. 22, p. 101108, Sept. 2021.
- [33] S. Mafi, Y. Abdelraziz, and R. Doczy, “Analysis of Gap Acceptance Behavior for Unprotected Right and Left Turning Maneuvers at Signalized Intersections using Data Mining Methods: A Driving Simulation Approach,” *Transp. Research Record*, vol. 2672, pp. 160–170, Dec. 2018.
- [34] Abhishek, M. A. A. Boon, and M. Mandjes, “Generalized gap acceptance models for unsignalized intersections,” *Math. Methods Oper. Research*, vol. 89, pp. 385–409, June 2019.
- [35] A. Zgonnikov, D. Abbink, and G. Markkula, “Should I stay or should I go? Evidence accumulation drives decision making in human drivers,” 2020. PsyArXiv.
- [36] M. Arafat, M. Hadi, T. Hunsan, and K. Amine, “Stop Sign Gap Assist Application in a Connected Vehicle Simulation Environment,” *Transp. Research Record*, vol. 2675, pp. 1127–1135, Sept. 2021.
- [37] Y. Hu, X. Jia, M. Tomizuka, and W. Zhan, “Causal-based Time Series Domain Generalization for Vehicle Intention Prediction,” in *IEEE Int. Conf. on Robotics Autom. (ICRA)*, pp. 7806–7813, May 2022.
- [38] E. Balal and R. L. Cheu, “Comparative Evaluation of Fuzzy Inference System, Support Vector Machine and Multilayer Feed-Forward Neural Network in Making Discretionary Lane Changing Decisions,” *Neural Netw. World*, vol. 28, no. 4, pp. 361–378, 2018.
- [39] D.-F. Xie, Z.-Z. Fang, B. Jia, and Z. He, “A data-driven lane-changing model based on deep learning,” *Transp. Research Part C: Emerg. Technol.*, vol. 106, pp. 41–60, Sept. 2019.
- [40] A. Das, M. N. Khan, and M. M. Ahmed, “Nonparametric Multivariate Adaptive Regression Splines Models for Investigating Lane-Changing Gap Acceptance Behavior Utilizing Strategic Highway Research Program 2 Naturalistic Driving Data,” *Transp. Research Record*, vol. 2674, pp. 223–238, May 2020.
- [41] B. Khelfa, I. Ba, and A. Tordeux, “Predicting highway lane-changing maneuvers: A benchmark analysis of machine and ensemble learning algorithms,” Apr. 2022. arXiv:2204.10807[physics].
- [42] A. Rudenko, L. Palmieri, M. Herman, K. M. Kitani, D. M. Gavrila, and K. O. Arras, “Human motion trajectory prediction: a survey,” *The Int. J. Robotics Research*, vol. 39, pp. 895–935, July 2020.
- [43] G. Lyons, J. Hunt, and F. McLeod, “A neural network model for enhanced operation of midblock signalled pedestrian crossings,” *Eur. J. Oper. Research*, vol. 129, pp. 346–354, Mar. 2001.
- [44] B. R. Kadali, P. Vedagiri, and N. Rathi, “Models for pedestrian gap acceptance behaviour analysis at unprotected mid-block crosswalks under mixed traffic conditions,” *Transp. Research Part F: Traffic Psychol. Behaviour*, vol. 32, pp. 114–126, July 2015.
- [45] R. Yao, W. Zeng, Y. Chen, and Z. He, “A deep learning framework for modelling left-turning vehicle behaviour considering diagonal-crossing motorcycle conflicts at mixed-flow intersections,” *Transp. Research Part C: Emerg. Technol.*, vol. 132, p. 103415, Nov. 2021.
- [46] R. Nagalla, P. Pothuganti, and D. S. Pawar, “Analyzing Gap Acceptance Behavior at Unsignalized Intersections Using Support Vector Machines, Decision Tree and Random Forests,” *Procedia Comput. Sci.*, vol. 109, pp. 474–481, Jan. 2017.
- [47] M. Yang, X. Wang, and M. Qudus, “Examining lane change gap acceptance, duration and impact using naturalistic driving data,” *Transp. Research Part C: Emerg. Technol.*, vol. 104, pp. 317–331, July 2019.
- [48] H. Müller, W. Müller, S. Marchand-Maillet, T. Pun, and D. M. Squire, “A Framework for Benchmarking in CBIR,” *Multimed. Tools Appl.*, vol. 21, pp. 55–73, Sept. 2003.
- [49] M. Zaffar, S. Garg, M. Milford, J. Kooij, D. Flynn, K. McDonald-Maier, and S. Ehsan, “VPR-Bench: An Open-Source Visual Place Recognition Evaluation Framework with Quantifiable Viewpoint and Appearance Change,” *Int. J. Comput. Vis.*, vol. 129, pp. 2136–2174, July 2021.
- [50] Y. Cao, C. Xiao, A. Anandkumar, D. Xu, and M. Pavone, “AdvDO: Realistic Adversarial Attacks for Trajectory Prediction,” Sept. 2022. arXiv:2209.08744 [cs].
- [51] R. Krajewski, J. Bock, L. Kloecker, and L. Eckstein, “The highD Dataset: A Drone Dataset of Naturalistic Vehicle Trajectories on German Highways for Validation of Highly Automated Driving Systems,” in *IEEE Int. Conf. on Intell. Transp. Syst. (ITSC)*, pp. 2118–2125, Nov. 2018.
- [52] R. Krajewski, T. Moers, J. Bock, L. Vater, and L. Eckstein, “The roundD Dataset: A Drone Dataset of Road User Trajectories at Roundabouts in Germany,” in *IEEE Int. Conf. on Intell. Transp. Syst. (ITSC)*, pp. 1–6, Sept. 2020.
- [53] E. Barnard and L. Wessels, “Extrapolation and interpolation in neural network classifiers,” *IEEE Control Syst. Mag.*, vol. 12, pp. 50–53, Oct. 1992.
- [54] G. Marcus, “The Next Decade in AI: Four Steps Towards Robust Artificial Intelligence,” Feb. 2020. arXiv: 2002.06177[cs].
- [55] M. van Bekkum, M. de Boer, F. van Harmelen, A. Meyer-Vitali, and A. t. Teije, “Modular design patterns for hybrid learning and reasoning systems,” *Appl. Intell.*, vol. 51, pp. 6528–6546, Sept. 2021.



Julian F. Schumann received the master's degree in 2021 in mechanical engineering from the TU Delft, The Netherlands, where since 2021, he is working toward the Ph.D. degree focusing on hybrid-AI models for behavior prediction of human traffic participants for use in automated vehicles. His research interests include modeling of human behavior, meaningful evaluation of such models, and the integration of knowledge-based and machine-learned models.



Jens Kober is an associate professor at the TU Delft, Netherlands. He worked as a postdoctoral scholar jointly at the CoR-Lab, Bielefeld University, Germany and at the Honda Research Institute Europe, Germany. He graduated in 2012 with a PhD Degree in Engineering from TU Darmstadt. For his research he received the 2018 IEEE RAS Early Academic Career Award and the 2022 RSS Early Career Award. His research interests include motor skill learning, imitation learning, interactive learning, and machine learning for control.



Arkady Zgonnikov received his Ph.D. degree from the University of Aizu, Japan, in 2014. Since then he worked as a postdoctoral researcher at the University of Galway, Ireland (funded by Irish Research Council) and Delft University of Technology, Netherlands (funded by the AiTech initiative). Since 2020, he has been an assistant professor at Delft University of Technology. His research interests include responsible AI and cognitive modeling, as well as applications of those to human-robot interactions in traffic and beyond.